

From a public address list to a costed sales pipeline.

Imagery, judgment, history, cost.

231,778 Knoxville addresses, a public list — no purchased leads, no CRM export. A vision model reads satellite and street-level imagery for every roof, cross-references storm history, and returns a scored, costed, reviewable pipeline inside a working app.

DELIVERED FOR A CLIENT — NOT MULE AI'S OWN PRODUCT

Delivered for a roofing company client — not named here, at their request. This was a technical pilot phase of that engagement: a vision model that turns a plain, public Knoxville address list into a qualified, costed sales pipeline a sales team could actually work from, not a Mule AI experiment run for its own sake.



Six stages. One address list.

Every property in the list passes through the same six stages — nothing is judged in isolation, and nothing skips the storm-history or confidence checks on the way to a score.

01

Ingest the address list

Started from a public address list covering Knoxville — 231,778 properties, no purchased leads, no CRM export.

02

Pull imagery per property

For each address, pulled satellite and street-level imagery to give the model an overhead and an eye-level view of the same roof.

03

Judge the roof with vision

A vision model rates each roof on material, age, complexity and visible damage signs — each dimension carrying its own confidence score rather than one blended guess.

04

Cross-reference storm history

Each property is checked against hail and wind event history for the area, so a plausible-looking roof with a storm in its recent past is not treated the same as one with a clean record.

05

Segment and cost the roof

The roof area is segmented from the imagery and used to estimate a replacement cost as a range, not a single number the model cannot actually back up.

06

Score, bucket and flag for review

Every property gets a 0–100 score, a bucket like "hot", and a needs-review flag wherever the model's own confidence is too low to act on without a person checking first.

Why this matters. A score with no confidence attached is just a guess wearing a number. Flagging the low-confidence reads is what keeps the other 44 from being judged the same way as the one the model wasn't sure about.

231,778

ADDRESSES
PROCESSED

45

RANKED PROSPECTS

4

JUDGMENT
DIMENSIONS

0–100

SCORE RANGE

6

PIPELINE STAGES

2

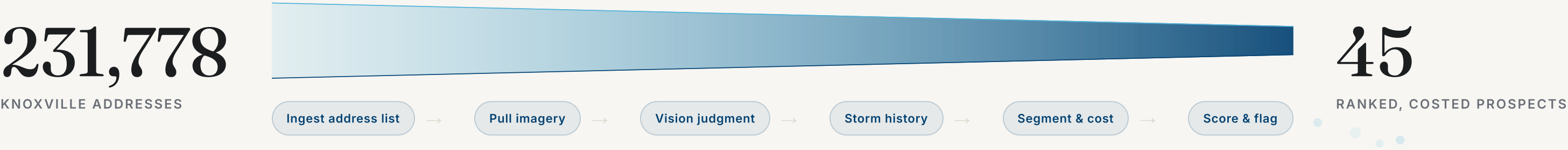
IMAGERY ANGLES
PER ROOF

0

LEADS PURCHASED

231,778 addresses. 45 prospects.

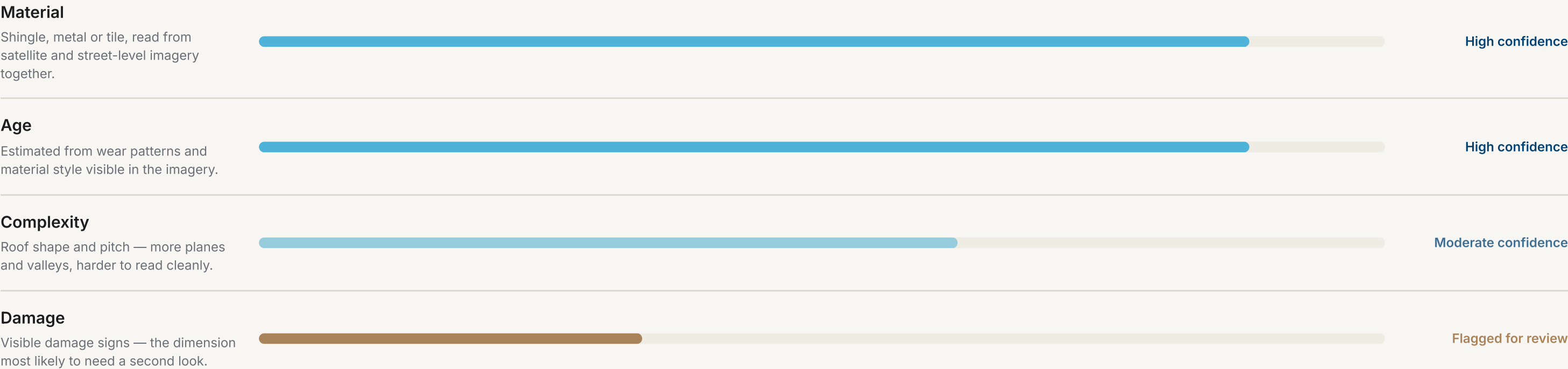
No property is scored or bucketed in the middle of this diagram without a number to show for it — every count here is one this pilot actually produced.



What this means for a business like yours. The input was a public list anyone can download — the kind of asset that sits in a shared drive doing nothing because nobody has a week to sort it by hand. This pipeline made it work without a person opening a single record.

Four dimensions. Four confidence scores.

The model never blends its read into one guess. Material, age, complexity and damage are scored separately, each carrying its own confidence — so a property can be judged precisely on three dimensions and flagged for a person on the fourth.

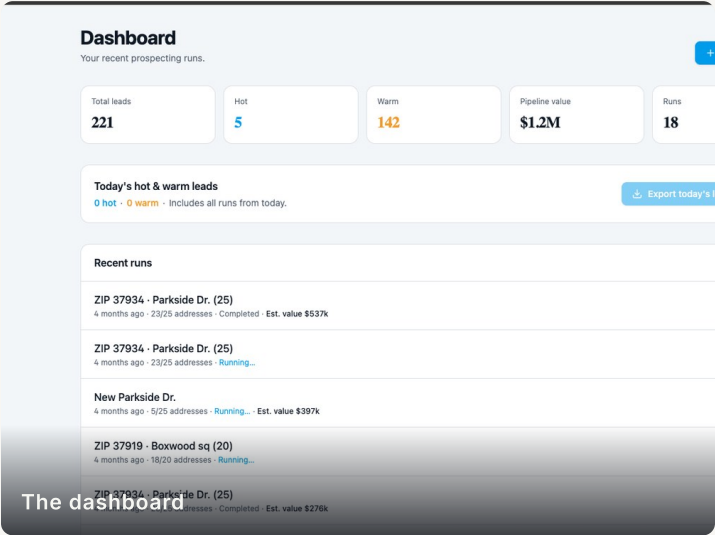


Illustrative — confidence is computed per property, not fixed by dimension. This shows the pattern, not one property’s actual reading.

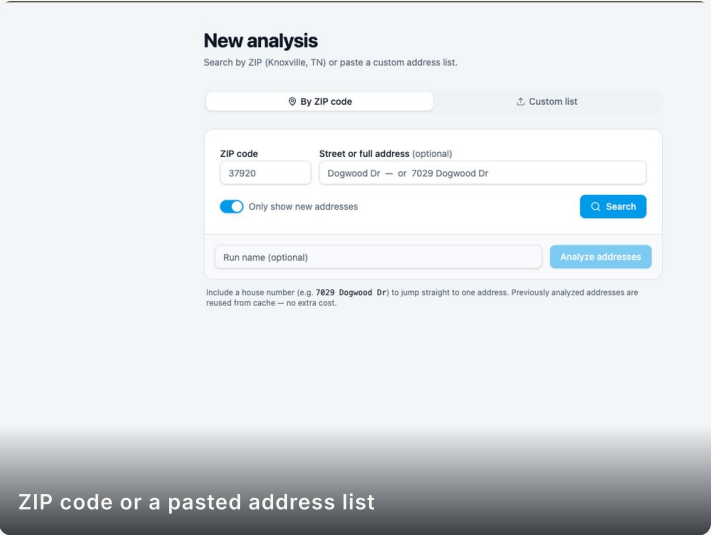
Segmenting and costing. The roof area is segmented from the same imagery and used to estimate a replacement cost as a range — not a single number the model cannot actually back up.

A working app, not a spreadsheet export.

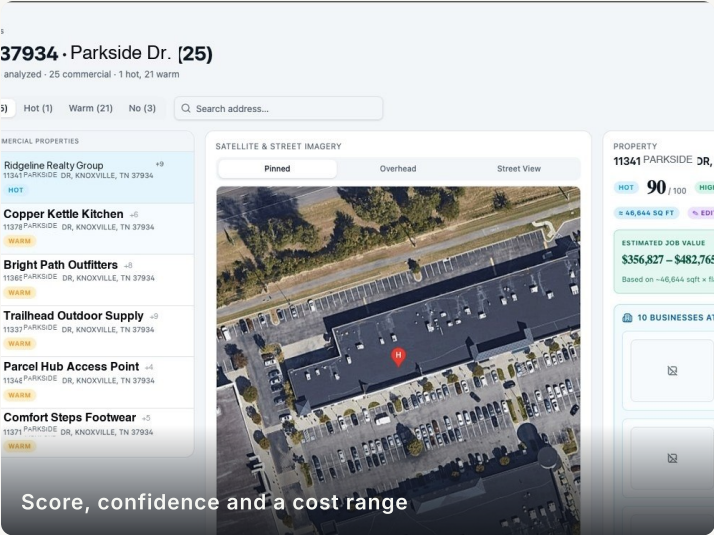
Every image here is the real pilot build — a run history, a scored property with its confidence and cost range, and the model’s own notes on one roof.



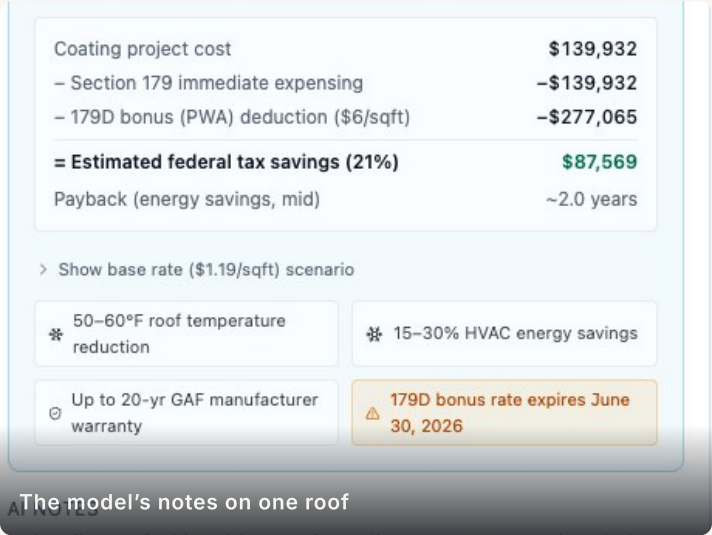
The dashboard



ZIP code or a pasted address list



Score, confidence and a cost range



The model’s notes on one roof

Every screen here is the real build delivered for this client, run over real Knoxville imagery — not a mockup and not the abstract diagram this deck used before the app had screens to show.

What the pipeline did not measure.

NOT A SALES CAMPAIGN

This was delivered as a technical pilot for the client, not a sales campaign Mule AI ran itself. What the client did with the 45 properties afterward — any calls made, any deals closed — was not instrumented by us, so there is no close rate, no revenue and no ROI to report, and none is claimed here. What was verified on our end is the pipeline itself: it runs end to end over a real address list and returns output a sales team could actually work from.

- What the client did with the 45 properties afterward was not tracked by Mule AI.
- No close rate.
- No revenue.
- No ROI — and none is claimed.
- Delivered as a technical pilot, not a sales campaign Mule AI ran itself.

“This was delivered as a technical pilot, not a sales campaign Mule AI ran itself.”

STATED PLAINLY, IN THIS CASE STUDY’S OWN HONESTY SECTION

231,778 addresses were never the hard part. Judging them was.

The input here was a public list anyone can download — the kind of asset that sits in a shared drive doing nothing because nobody has a week to sort it by hand. The pipeline turned it into 45 ranked, costed, reviewable prospects without a person opening a single record. That is the shape of the audit’s "revenue opportunity" finding: not a new tool bolted onto the business, but an existing pile of data made to work for the first time.



01	02	03	04	05	06
Ingest the address list	Pull imagery per property	Judge the roof with vision	Cross-reference storm history	Segment and cost the roof	Score, bucket and flag for review
Started from a public address list covering Knoxville — 231,778 properties, no purchased leads, no CRM export.	For each address, pulled satellite and street-level imagery to give the model an overhead and an eye-level view of the same roof.	A vision model rates each roof on material, age, complexity and visible damage signs — each dimension carrying its own confidence score rather than one blended guess.	Each property is checked against hail and wind event history for the area, so a plausible-looking roof with a storm in its recent past is not treated the same as one with a clean record.	The roof area is segmented from the imagery and used to estimate a replacement cost as a range, not a single number the model cannot actually back up.	Every property gets a 0–100 score, a bucket like "hot", and a needs-review flag wherever the model’s own confidence is too low to act on without a person checking first.